



Contract Design With Safety Inspections

ALIREZA FALLAH, University of California, Berkeley, USA

MICHAEL I. JORDAN, University of California, Berkeley, USA

We study the role of regulatory inspections in a contract design problem in which a principal interacts separately with multiple agents. Each agent's hidden action includes a dimension that determines whether they undertake an extra costly step to adhere to safety protocols. The principal's objective is to use payments combined with a limited budget for random inspections to incentivize agents towards safety-compliant actions that maximize the principal's utility. We first focus on the single-agent setting with linear contracts and present an efficient algorithm that characterizes the optimal linear contract, which includes both payment and random inspection. We further investigate how the optimal contract changes as the inspection cost or the cost of adhering to safety protocols vary. Notably, we demonstrate that the agent's compensation increases if either of these costs escalates. However, while the probability of inspection decreases with rising inspection costs, it demonstrates nonmonotonic behavior as a function of the safety action costs. Lastly, we explore the multi-agent setting, where the principal's challenge is to determine the best distribution of inspection budgets among all agents. We propose an efficient approach based on dynamic programming to find an approximately optimal allocation of inspection budget across contracts. We also design a random sequential scheme to determine the inspector's assignments, ensuring each agent is inspected at most once and at the desired probability. Finally, we present a case study illustrating that a mere difference in the cost of inspection across various agents can drive the principal's decision to forego inspecting a significant fraction of them, concentrating its entire budget on those that are less costly to inspect.

CCS Concepts: • **Theory of computation** → **Algorithmic game theory and mechanism design**; • **Applied computing** → **Economics**.

Additional Key Words and Phrases: Contract Design, Inspection, Safety

ACM Reference Format:

Alireza Fallah and Michael I. Jordan. 2024. Contract Design With Safety Inspections. In *The 25th ACM Conference on Economics and Computation (EC '24)*, July 8–11, 2024, New Haven, CT, USA. ACM, New York, NY, USA, 23 pages. <https://doi.org/10.1145/3670865.3673489>

1 Introduction

The rapid growth of data-oriented applications has led to a surge in new products and services offered by companies that provide personalized user experiences. Alongside the benefits of personalization, however, there are growing concerns about potential unwanted side effects that are not necessarily revealed or declared by these companies. For example, many users wonder whether their data is stored securely every time they enter their sensitive information on an online platform, especially given that instances of data breaches and cyberattacks have become common.

Authors' Contact Information: Alireza Fallah, afallah@berkeley.edu, University of California, Berkeley, Berkeley, CA, USA; Michael I. Jordan, jordan@cs.berkeley.edu, University of California, Berkeley, Berkeley, CA, USA.



This work is licensed under a Creative Commons Attribution International 4.0 License.

EC '24, July 8–11, 2024, New Haven, CT, USA

© 2024 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-0704-9/24/07

<https://doi.org/10.1145/3670865.3673489>

In the realm of healthcare, the failure of companies to reveal negative and undesired side effects of drugs has had devastating consequences, as in the Vioxx case in the early 2000s and the opioid epidemic exacerbated by Purdue Pharma in the United States. These episodes demonstrate how inadequate disclosure and transparency can have severe consequences for public health.

Another concerning example pertains to the design practices of leading tech companies. There are valid concerns about whether they pay enough attention to safety, security, and reliability when creating new products and algorithms. Releasing cutting-edge technologies without adequate testing or safeguards raises questions about potential risks to users and the wider public.

Beyond simply documenting such concerns, the question arises as to how we might incentivize platforms to follow safety and security measures. In this work, we propose a mechanism design framework based on contract theory to provide such incentives, focusing on an incentive-producing role for *inspections*. Building on the classical principal-agent model, where the agents (in this case, the platforms) take hidden and costly actions that result in a reward for the principal, we introduce an additional dimension to the model. Specifically, each agent must also decide whether to take a costly step to adhere to safety and security measures or disregard them. In the event that the agent neglects these measures, negative side effects may occur with a non-negligible probability, leading to adverse consequences for the principal.

The principal offers compensation to each agent based on the reward generated from the principal's action. Moreover, the principal retains the option to conduct a random and costly inspection, revealing whether the agent has complied with the safety and security measures without disclosing any information about the agent's underlying action. In addition, although the principal maintains a separate contract with each agent, the collective design of these contracts is constrained by the principal's limited inspection budget; in particular the number of inspectors. In essence, the principal's objective is to identify the optimal set of contracts, consisting of payments and inspections, that maximizes its overall utility subject to its inspection budget.

Our work introduces a framework to model and analyze the role of regulatory and partial inspections in contract design. These inspections aim to confirm the agent's compliance with laws emphasizing societal considerations, such as safety. Our first set of results focuses on the design of linear contracts in a single-agent setting. We characterize the optimal probability of inspection as a function of the ratio of the reward paid to the agent, and show that it is a piecewise convex and decreasing function. Our detailed characterization enables us to design an algorithm that finds the optimal linear contract in quasi-linear time in the number of actions.

Furthermore, we determine that when the inspection becomes more costly for the principal, they tend to reduce the inspection probability but compensate the agent more to ensure they are motivated to adhere to safety regulations. Similarly, if the cost associated with observing safety protocols rises for the agent, the principal modifies the optimal contract by increasing the agent's compensation to guarantee adherence to safe practices. Surprisingly, in such scenarios, the optimal probability of inspection does not necessarily increase. In fact, it turns out that the increase in payment might even allow the principal to decrease the inspection level. This observation underscores that merely ramping up regulatory inspection is not always the most effective strategy, especially in contexts where safety comes at a high cost.

Next, we turn our attention to a multi-contract setting with multiple agents. Here, due to their budget constraints, the principal cannot assign the optimal inspection level to each agent's contract and must thus determine the best allocation of inspections. We demonstrate that the principal's utility from each contract, when considered as a function of the maximum permissible inspection, is a piecewise concave and weakly increasing function. Further, we establish that the principal's problem in this context is closely related to the multiple-choice knapsack problem. We then introduce a dynamic-programming-based algorithm that finds an ϵ -approximate solution, with a time complexity that is polynomial in terms of the number of agents, the number of actions, and $1/\epsilon$. We further use a random procedure to assign inspectors to agents, ensuring that each agent is inspected with the targeted probability of inspection. It's important to note that we cannot determine each inspector's schedule independently since each agent must be inspected no more than once. Accordingly, our procedure determines the assignment of inspectors sequentially. For every pair of consecutive inspectors, there is at most one common agent they might inspect. Nevertheless, we design each inspector's assignment distribution based on the preceding inspector to prevent overlaps and ensure that the desired inspection level is attained.

Finally, we present a case study illustrating an intriguing dynamic: even when dealing with identical agents who solely differ in terms of the cost incurred by the principal to inspect them, the principal may decide to abstain from inspecting a substantial proportion of those with higher inspection costs. This observation suggests that agents have the potential leverage to influence the principal's action by increasing the cost barriers to monitor their adherence to safety protocols. In fact, by doing so, they can dissuade the principal from monitoring them and hence increase their welfare.

1.1 Related Work

Our model builds on the hidden-action principal-agent model [Grossman and Hart, 1992]. Within this general framework, the study of costly state verification dates back to the work of Gale and Hellwig [1985], Townsend [1979] for debt contracts (see also chapter 5.3 in [Bolton and Dewatripont, 2004] for discussion on costly verification or disclosure).

Our paper aligns most closely with the literature on contract design with random monitoring [Barbos, 2022, Jost, 1991, 1996, Strausz, 1997]. In [Jost, 1991] the principal decides randomly to monitor the agent's action through a costly inspection, and it is argued that in the optimal contract, the principal either performs the inspection or pays the inspection cost to the agent. The work in [Jost, 1996] considers a model where the principal has private information regarding their monitoring costs. In [Barbos, 2022], the inspection uncovers the agent's exact action with a certain probability, and in other cases, provides no information. Our work departs from the existing literature in three primary respects: First, we introduce a model of partial inspection in contract theory. In our framework, the inspection fully discloses adherence to safety standards but reveals no information about the other dimension of the agent's actions. Second, we establish the computational complexity of both equilibrium characterization and its comparative statics. And third, our results span both single-agent and multi-agent settings.

It is also worth mentioning that while our work focuses on one-period contract design, the optimal randomized inspection in dynamic contracts over time has been considered in the literature as well [Ball and Knoepfle, 2023, Chen et al., 2020, Orlov, 2022, Varas et al., 2020].

In mechanism design, our paper relates to the literature on mechanism design with costly inspection or verification [Ben-Porath et al., 2014, Li, 2020, Mylovanov and Zapechelnyuk,

2017]. The focus in this line of work is on the principal’s problem of allocating an object based on the reports from the agents of their private types, reports which are subject to potential inspection by the principal. These papers along with others such as [Erlanson and Kleiner, 2020, Halac and Yared, 2020] use inspection as a tool when monetary payments are not feasible; our work, on the other hand, allows for both payments and inspections. This brings our work closer to [Alaei et al., 2020] in the mechanism design literature, where the assumption is that an auctioneer can use both payments and a (full yet deferred) inspection.

Our work is also related to the literature on partial or probabilistic verification in mechanism design [Ball and Kattwinkel, 2019, Caragiannis et al., 2012, Ferraioli and Ventre, 2018, Green and Laffont, 1986]. In these models, an agent has a private type, and inspections can differentiate certain type pairs from each other but might be ineffective with others. Our safety inspection model can also be viewed in this model: pairs comprising a safe action and an unsafe action are discernible, whereas pairs consisting of two safe or two unsafe actions remain indistinct. In these works, however, inspections are presumed to be cost-free, and their primary focus is to characterize the class of social choice functions that can be implemented truthfully.

Our motivations are similar to those in the literature on regulatory inspections for incentivizing companies to adhere to standard policies [Choe and Fraser, 1999, Ferraro, 2008, Harrington, 1988]. For instance, in [Harrington, 1988], companies are partitioned into two groups. One group is inspected more frequently, and its members face steeper fines if found to be violating protocols. Adhering to standard protocols incurs a cost, and firms can be shifted from one group to another—either via a reward or a punishment—based on their performance.

Our work also relates to the literature on computational aspects of contract theory [Babaioff et al., 2006, Dütting et al., 2022, Dütting et al., 2021, Zhu et al., 2023]. As in this line of work, we focus on the class of linear contracts, given their simplicity and interpretability, and the fact that linear contracts have been shown to be robust to the unknown actions [Carroll, 2015] or unknown distributions [Dütting et al., 2019].

The remainder of the paper is organized as follows: Section 2 introduces the model and setting considered in this study. Section 3 examines the single-agent scenario, characterizing the optimal linear contract with partial safety inspection and exploring certain comparative statics. In Section 4, we expand our analysis to a multi-contract context with a limited inspection budget, drawing connections to the multiple-choice knapsack problem which helps us in developing an algorithm for finding the optimal contract. We conclude the paper in Section 5, and Appendix A contains proofs omitted from the main text.

2 The Model

We consider a multi-agent setting with m agents and one principal. For agent $\ell \in [m] := \{1, \dots, m\}$, we use the notation (a^ℓ, s^ℓ) to represent the agent’s action, where $a^\ell \in \mathcal{A}^\ell = \{a_1^\ell, \dots, a_n^\ell\}$ determines the effort the agent invests in providing its service,¹ and $s^\ell \in \{0, 1\}$ indicates whether the agent considers safety measures. Accordingly, we call actions with $s^\ell = 1$ and $s^\ell = 0$ as *safe* and *unsafe* actions, respectively. For any $i \in [n]$, the cost of action a_i^ℓ is denoted by $c_i^\ell \geq 0$. Additionally, the cost of complying with safety measures (i.e., $s^\ell = 1$) is denoted by $\kappa_S^\ell \geq 0$. Hence, the total cost of action (a_i^ℓ, s^ℓ) is given by $c_i^\ell + \mathbb{1}_{s^\ell=1} \kappa_S^\ell$.

¹Here, to simplify the notation, we assume all agents have n actions. However, our analysis will remain the same when they have different numbers of actions.

When agent ℓ takes action (a^ℓ, s^ℓ) , a random reward $r^\ell \in \mathcal{R}^\ell \subseteq \mathbb{R}^{\geq 0} \cup \{-\infty\}$ accrues to the principal. Specifically, if the agent takes action $(a^\ell, 1)$ which includes adherence to safety protocols, a nonnegative reward r^ℓ is generated with probability $f^\ell(r^\ell|a^\ell)$. Conversely, if the agent takes action $(a^\ell, 0)$, neglecting safety measures, there is a probability α^ℓ that negative side effects occur, leading to a reward of $-\infty$. With the complementary probability, these side effects do not materialize, and a nonnegative reward r^ℓ is generated with probability $f^\ell(r^\ell|a^\ell)$, similar to the case where the agent followed the safety measures.

The contract between each agent and the principal comprises two elements. The principal, upon observing the reward r^ℓ from agent ℓ , compensates them for their action by paying them $t^\ell(r^\ell) \geq 0$ with the condition $t^\ell(-\infty) = 0$. In other words, the principal pays nothing if side effects occur. Additionally, the principal has the option to perform an inspection, with probability β^ℓ and at a cost of κ_i^ℓ , before the reward realization, which reveals whether the agent adhered to safety measures or not, i.e., it reveals the value of s^ℓ . What connects all of the contracts is the fact that the principal has a limited capacity for only $B \in \mathbb{Z}^+$ unit of inspection, meaning that the condition $\sum_{\ell=1}^m \beta^\ell \leq B$ should hold (one can think of this as having B inspectors available). We also assume that the two events of inspection and occurrence of side effects are independent.

We denote the expected reward to the principal and the expected payment to the ℓ -th agent as a result of agent ℓ taking action $(a_i^\ell, 1)$ by R_i^ℓ and T_i^ℓ , respectively, i.e.,

$$R_i^\ell := \mathbb{E}_{r^\ell \sim f^\ell(\cdot|a_i^\ell)}[r] \quad \text{and} \quad T_i^\ell := \mathbb{E}_{r^\ell \sim f^\ell(\cdot|a_i^\ell)}[t^\ell(r)]. \quad (1)$$

2.1 Implementable Actions

We denote agent ℓ 's expected utility when taking action (a_i^ℓ, s^ℓ) by $\mathcal{U}_a^\ell(a_i^\ell, s^\ell)$, which is defined as follows:

$$\mathcal{U}_a^\ell(a_i^\ell, s^\ell) = \begin{cases} T_i^\ell - c_i^\ell - \kappa_S^\ell & s^\ell = 1, \\ (1 - \beta^\ell)(1 - \alpha^\ell)T_i^\ell - c_i^\ell & s^\ell = 0. \end{cases} \quad (2)$$

We say an action (a^ℓ, s^ℓ) is *implementable* for agent ℓ if there exists a contract, consisting of $(t^\ell(\cdot), \beta^\ell)$, such that the following two conditions hold:

- *Incentive compatibility (IC)*: the agent has no incentive to deviate and choose another action, i.e., $\mathcal{U}_a^\ell(a^\ell, s^\ell) \geq \mathcal{U}_a^\ell(a', s')$, for any other $a' \in \mathcal{A}^\ell$ and $s' \in \{0, 1\}$.
- *Individual rationality (IR)*: the agent is not better off by not taking the contract at all, i.e., $\mathcal{U}_a^\ell(a^\ell, s^\ell) \geq 0$.

In other words, IC and IR together ensure that, the best response of the agent to the offered contract is to take the action (a^ℓ, s^ℓ) . We let $\mathcal{I}^\ell$ denote the set of implementable actions for agent ℓ .

It is worth noting that having $(a^\ell, s^\ell) \in \mathcal{I}^\ell$ for all $\ell \in [m]$ does not necessarily imply that the tuple of actions $((a^\ell, s^\ell))_{\ell=1}^m$ is implementable for all agents simultaneously, as we have not yet taken into account the constraint $\sum_{\ell=1}^m \beta^\ell \leq B$. To this end, we also say a tuple of actions $((a^\ell, s^\ell))_{\ell=1}^m$ is *fully implementable* (and denote the set of such tuples by $\mathcal{I}$) if, for any ℓ , (a^ℓ, s^ℓ) is implementable by some contract $(t^\ell(\cdot), \beta^\ell)$ such that $\sum_{\ell=1}^m \beta^\ell \leq B$.

2.2 The Principal's Problem

Let us denote the principal's expected utility from agent ℓ when that agent takes action (a_i^ℓ, s^ℓ) by $\mathcal{U}_p^\ell(a_i^\ell, s^\ell)$. For $s^\ell = 0$, this utility is $-\infty$ as the side effects arise with a non-zero probability and lead to a reward of $-\infty$. Hence, the principal would strictly prefer safe

actions over unsafe ones. For the case of $s^\ell = 1$, the expected utility of principal is given by

$$\mathcal{U}_p^\ell(a_i^\ell, 1) = R_i^\ell - T_i^\ell - \beta^\ell \kappa_i^\ell. \quad (3)$$

In addition, we denote the total expected utility of the principal when agents take actions $((a_i^\ell, s_i^\ell))_{\ell=1}^m$ by $\mathcal{U}_p(((a_i^\ell, s_i^\ell))_{\ell=1}^m)$, defined as follows:

$$\mathcal{U}_p(((a_i^\ell, s_i^\ell))_{\ell=1}^m) = \begin{cases} \sum_{\ell=1}^m \mathcal{U}_p^\ell(a_i^\ell, 1) & \text{if } s_i^\ell = 1 \text{ for all } \ell, \\ -\infty & \text{otherwise.} \end{cases} \quad (4)$$

Now, the principal's problem can be seen as designing a contract that incentivizes agents to play the tuple of actions that maximizes her expected utility among all implementable tuples of actions. In other words, the principal's problem can be cast as finding the contract that implements the solution to the following maximization problem:

$$\begin{aligned} \max_{((a_i^\ell, s_i^\ell))_{\ell=1}^m} & \mathcal{U}_p(((a_i^\ell, s_i^\ell))_{\ell=1}^m) \\ \text{s.t.} & ((a_i^\ell, s_i^\ell))_{\ell=1}^m \in \mathcal{I}. \end{aligned} \quad (5)$$

We make the following assumptions throughout the paper:

ASSUMPTION 1. *For every agent, we assume different actions have different costs, and moreover, an action with a higher cost also has a higher expected reward. Also, we assume that no two actions of any given agent have the same cost.*

Assumption 1 ensures that no action dominates another one, meaning that it leads to a higher expected reward at a lower cost. Under this assumption, and without loss of generality, we assume $c_1^\ell < c_2^\ell < \dots < c_n^\ell$ and $R_1^\ell < R_2^\ell < \dots < R_n^\ell$ for every $\ell \in [m]$.

ASSUMPTION 2. *For every agent ℓ , we have $\max_i (R_i^\ell - c_i^\ell) > \kappa_S^\ell$.*

Note that if this assumption does not hold for an agent, it implies that, even if that agent receives the entire reward as payment, their utility would remain nonpositive for any safe action. In simpler terms, this assumption guarantees there is a way to make at least one safe action implementable for every agent.

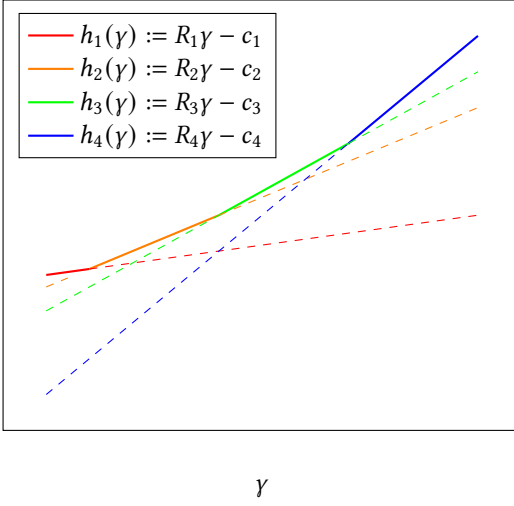
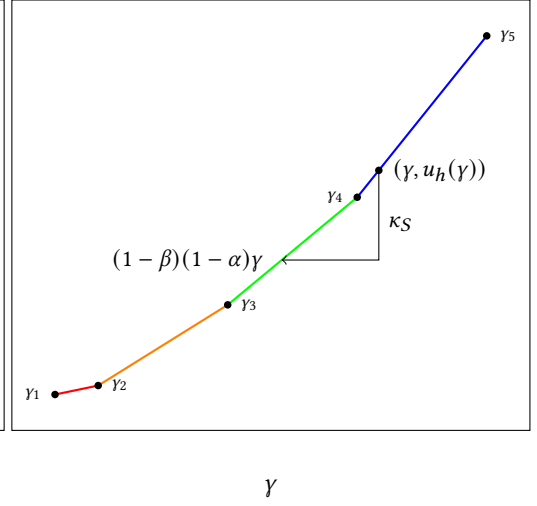
Our focus is on linear contracts, which are contracts for which the agent's payment is directly proportional to the reward received by the principal. General contracts can often be complex and challenging to understand or implement, while linear contracts offer a simpler and more practical alternative. Moreover, under reasonable assumptions, linear contracts are known to be worst-case optimal [Dütting et al., 2019]. In our setting, focusing on linear contracts allow us to better isolate the role of partial inspection in contract design over a useful and intuitive class of contracts.

More formally, we consider payment function $t^\ell(r) = \gamma^\ell r$, for some $\gamma^\ell \in [0, 1]$ chosen by the principal. Hence, the principal has two sets of parameters to choose: γ^ℓ , the ratio of the reward paid to agent ℓ , and β^ℓ , the probability of performing inspection regarding the safety of agent ℓ 's action. Note also that in this case, $T_i^\ell = \gamma^\ell R_i^\ell$.

3 The Single-Agent Setting

To gain a better understanding of the nature of the problem, we begin with the single-agent setting, i.e., $m = 1$. Without loss of generality, we assume $B = 1$ in this case. To simplify the notation, we drop the superscripts throughout this section.

Let us first consider what happens if the principal is not allowed to do the inspection, i.e., β is set to 0. In this case, the principal should intuitively offer a higher payment ratio to

Fig. 1. Illustration of $u_h(\cdot)$ Fig. 2. How to characterize $\beta(\gamma)$

persuade the agent to adhere to safety protocols. However, as the following result shows, this may not be enough.

LEMMA 1. *If $\alpha < \frac{\kappa_S}{R_n}$, then there is no safe action that is implementable by a linear contract without inspection.*

This result shows that, when the negative side effects are rare enough, the principal would need to perform the inspection to keep the agent committed to the safety measures. Otherwise, when the occurrence probability of side effect is small enough, the agent takes a chance in not abiding the safety protocols. Next, we continue by characterizing the properties of the linear contract. By IC constraint, if a linear contract (γ, β) implements the safe action $(a_i, 1)$, then $\mathcal{U}_a(a_i, 1) \geq \mathcal{U}_a(a_j, 1)$ for any $j \neq i$. This simplifies to $\gamma R_i - c_i \geq \gamma R_j - c_j$ for any $j \neq i$.

Now, inspired by [Dütting et al., 2019], we develop the following geometric characterization: for any $i \in [n]$, let us define the linear function $h_i : [0, 1] \rightarrow \mathbb{R}$ as $h_i(\gamma) = \gamma R_i - c_i$. As we stated above, if the action $(a_i, 1)$ is implementable by (γ, β) , then $h_i(\gamma) \geq h_j(\gamma)$ for any $j \neq i$. As a result, to find the set of implementable actions, we need to characterize the upper envelope of the set of functions $\{h_j(\cdot)\}_{j=1}^n$, denoted by $u_h(\cdot)$, which is a piecewise linear and increasing function (see Figure 1 for an example). Each segment of $u_h(\cdot)$ corresponds to a specific $h_i(\cdot)$, representing the dominant function within that segment. This implies that, in that segment, $(a_i, 1)$ stands as the sole implementable safe action (for the γ values pertaining to that segment), with no incentive to divert to an alternative *safe* action. Note that, to satisfy the IC constraint, we must also ensure that there are no profitable deviations to unsafe actions. Additionally, the IR constraint needs to be examined as well.

Next, using this derivation, we establish the following result on implementable safe actions and their corresponding linear contracts:

PROPOSITION 1. *Suppose Assumptions 1 and 2 hold. There exist $0 = \gamma_0 < \gamma_1 < \gamma_2 < \dots < \gamma_k < \gamma_{k+1} = 1$ such that the following holds:*

- (1) *No safe action is implementable for $\gamma < \gamma_1$.*
- (2) *There exists a set of actions $i_1 < \dots < i_k$ such that, for any $j \in [k]$, the action $(a_{i_j}, 1)$ is implementable with $\gamma \in [\gamma_j, \gamma_{j+1}]$. Moreover, $(a_{i_j}, 1)$ is the only implementable safe action for $\gamma \in (\gamma_j, \gamma_{j+1})$.*
- (3) *For any $\gamma \geq \gamma_1$, there exists $\beta(\gamma)$ such that (γ, β) implements a safe action if and only if $\beta \geq \beta(\gamma)$.*

PROOF. First, note that, in order for a safe action $(a, 1)$ to be implementable, IR should hold as well. Hence, if $u_h(\gamma) \leq \kappa_S$, no safe action would be implementable by γ . As a result, no safe action is implementable for γ below

$$\gamma_1 := (u_h)^{-1}(\kappa_S). \quad (6)$$

Also, note that $u_h(1) = \max_j (R_j - c_j)$, which by assumption is greater than κ_S . Therefore, $\gamma_1 < 1$.

Now, let us focus on $\gamma \geq \gamma_1$. As we stated earlier, $u_h(\cdot)$ is a piecewise linear and increasing function. Hence, there exist $i_1 < \dots < i_k$ and $\gamma_1 < \gamma_2 < \dots < \gamma_k < \gamma_{k+1} = 1$ such that, for any $j \in [k]$, we have

$$u_h(\gamma) = h_{i_j}(\gamma) \text{ for any } \gamma \in [\gamma_j, \gamma_{j+1}]. \quad (7)$$

In other words, on segment $[\gamma_j, \gamma_{j+1}]$, $u_h(\cdot)$ is equal to $h_{i_j}(\cdot)$. Hence, for $\gamma \in [\gamma_j, \gamma_{j+1}]$ the only safe action that can potentially be implemented is $(a_{i_j}, 1)$. What remains is to rule out deviation to unsafe actions. Recall that the agent's utility from an unsafe action $(a_v, 0)$ is given by

$$(1 - \beta)(1 - \alpha)\gamma R_v - c_v \quad (8)$$

which is, in fact, $h_v((1 - \beta)(1 - \alpha)\gamma)$. As a consequence, the maximum utility that an agent can obtain from an unsafe action is given by $\max_v h_v((1 - \beta)(1 - \alpha)\gamma)$ which is $u_h((1 - \beta)(1 - \alpha)\gamma)$. Now, the IC constraint would require us to have

$$u_h(\gamma) - \kappa_S \geq u_h((1 - \beta)(1 - \alpha)\gamma). \quad (9)$$

Notice that the left-hand side is nonnegative, since $\gamma \geq \gamma_1$. Also, u_h is an increasing function, which starts from a negative value, i.e., $u_h(0) = \max_j (-c_j)$. Hence, we can choose β large enough such that (9) holds (see Figure 2 for an illustration). $\square$

Note that, for any $j \in [k]$, any $\gamma \in [\gamma_j, \gamma_{j+1}]$ with $\beta \geq \beta(\gamma)$ makes action $(a_{i_j}, 1)$ implementable for the agent. Since increasing β would only decrease the principal's utility, without loss of generality, we could assume the principal picks $\beta = \beta(\gamma)$. Next, we characterize how this probability of inspection $\beta(\gamma)$ changes as a function of γ .

LEMMA 2. *Suppose Assumptions 1 and 2 hold, and recall the definition of $\beta(\gamma)$ from Proposition 1. Then, $\beta(\gamma)$ is a decreasing function of γ . Moreover, it is strictly decreasing when $\beta(\gamma) > 0$.*

This result formalizes an intuitive observation: if the principal aims to reduce agent's payment, they should increase the inspection probability; conversely, if the principal wishes to avoid expensive inspections, they should offer higher compensation to the agent, encouraging adherence to safety protocols.

A natural question arises at this point: how should the principal determine the optimal trade-off between the agent's payment and the cost of inspection? Let γ fall within the interval $[\gamma_j, \gamma_{j+1}]$ for some j . Recall that, in this case, the principal's utility is given by $(1 - \gamma)R_{i_j} - \beta(\gamma)\kappa_I$. Consequently, the marginal cost associated with increasing the agent's

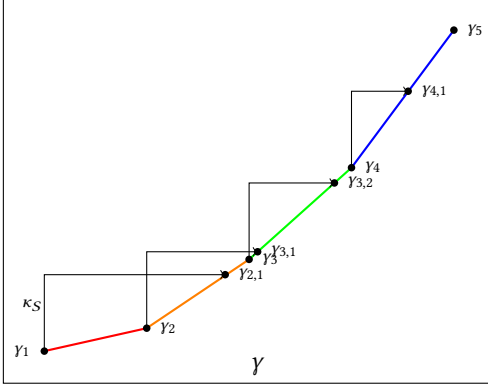


Fig. 3. How $\{\{\gamma_{j,q}\}_{q=1}^k\}_{j=1}^k$'s are defined

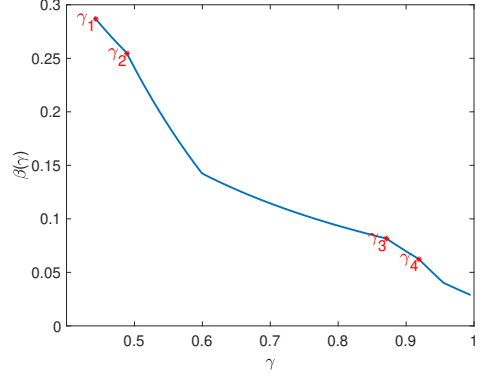


Fig. 4. $\beta(\gamma)$ as a function of γ (see Remark 1 for the details)

payment is $R_{i,j}$, while the marginal cost of inspection is κ_I . The next result helps us to find the appropriate γ that balances this trade-off.

LEMMA 3. *Under the premise of Proposition 1, and for any $j \in [k]$, $\beta(\gamma)$ is a convex function over the interval $[\gamma_j, \gamma_{j+1}]$. Moreover, it is strictly convex when $\beta(\gamma) > 0$.*

Proof sketch: Recall that, as illustrated in Figure 2, $\beta(\gamma)$ is given by

$$\beta(\gamma) = \max \left\{ 1 - \frac{u_h^{-1}(u_h(\gamma) - \kappa_S)}{\gamma(1 - \alpha)}, 0 \right\}. \quad (10)$$

Now, let us determine where $\beta(\cdot)$ could be nondifferentiable. As γ sweeps over the interval $[\gamma_j, \gamma_{j+1}]$, the corresponding $\tilde{\gamma} := u_h^{-1}(u_h(\gamma) - \kappa_S)$ may fall in this segment $[\gamma_j, \gamma_{j+1}]$ or one of the previous segments $[\gamma_i, \gamma_{i+1}]$ for some $i < j$. For instance, in the example illustrated in Figure 2, for γ marked on the plot, $\tilde{\gamma}$ falls within the previous segment. Hence, there exists a sequence $\gamma_j = \gamma_{j,0} < \gamma_{j,1} < \dots < \gamma_{j,v_j} < \gamma_{j,v_j+1} = \gamma_{j+1}$ such that $\tilde{\gamma}$ falls within the same segment for any $\gamma \in (\gamma_{j,q}, \gamma_{j,q+1})$. See Figure 3 for an illustration on how the $\{\gamma_{j,q}\}$ are defined.

Now, for any q , the $\beta(\cdot)$ function is differentiable over the interval $(\gamma_{j,q}, \gamma_{j,q+1})$. To prove Lemma 3, we first establish that $\beta(\cdot)$ is indeed convex over each interval $(\gamma_{j,q}, \gamma_{j,q+1})$ by showing that its derivative is increasing there. Finally, we present an argument detailing how the convexity of $\beta(\cdot)$ over the entire interval $[\gamma_j, \gamma_{j+1}]$ can be inferred from its derivative across these subintervals. $\square$

REMARK 1. *It is worth noting that while $\beta(\cdot)$ is a convex function over each interval $[\gamma_j, \gamma_{j+1}]$, it is not necessarily convex over the whole interval $[\gamma_1, 1]$. See Figure 4 for an example with $n = 6$ actions with the following parameters:*

$$[R_i]_{i=1}^6 = [2, 3, 7, 9, 11, 13], \quad [c_i]_{i=1}^6 = [1, 1.2, 2.1, 3.1, 4.8, 6.6], \quad \text{and } \kappa_I = \kappa_S = 1. \quad (11)$$

Now, having the characterization of the set of implementable actions, we focus on finding the optimal contract which maximizes the principal's utility. We start by showing that it can be found efficiently.

THEOREM 1. *Suppose Assumptions 1 and 2 hold. Then, the optimal linear contract (γ^*, β^*) can be characterized in time $O(n \log n)$.*

ALGORITHM 1: Computing the optimal contract in the single-agent setting**Input:** The reward and cost of different actions $(R_i, c_i)_{i=1}^n$ Find the upper envelope function $u_h(\cdot)$ by finding the convex hull of $(R_i, c_i)_{i=1}^n$ using the Graham Scan Algorithm [Graham, 1972];Denote the convex hull by $(R_{i_j}, c_{i_j})_{j=1}^k$ such that $i_1 < \dots < i_k$;Find $\gamma_1 < \dots < \gamma_k$ as defined in Proposition 1 and illustrated in Figure 2.Find all the nondifferentiable points $\{\{\gamma_{j,q}\}_{q=0}^{v_j+1}\}_{j=1}^k$ as illustrated in Figure 3;Let S denote the set of contracts corresponding to $\{\{\gamma_{j,q}\}_{q=0}^{v_j+1}\}_{j=1}^k$ and the associated utility of the principal;**for** $j = 1$ **to** k **do** **for** $q = 1$ **to** v_j **do** If the following equation has a solution over $[\gamma_{j,q}, \gamma_{j,q+1}]$, then add it to S :

$$\frac{c_{i_j} - c_{i_{jq}} + \kappa_S}{R_{i_{jq}} \gamma^2 (1 - \alpha)} = \frac{R_{i_j}}{\kappa_I};$$

end**end****Output:** Pick the contract(s) in S that maximizes the principal's utility

PROOF. The first step is to characterize u_h . Notice that, using a duality argument, we could transfer the problem of finding the upper envelope function to the problem of finding the convex hull of the set of points $(R_i, c_i)_{i=1}^n$, which can be done in $O(n \log n)$ using the Graham Scan Algorithm [Graham, 1972]. This allows us to find the set $\{i_1, \dots, i_k\}$ as defined in Proposition 1, and hence, the points $\gamma_1, \dots, \gamma_k$ in time $O(k \log k)$ which is bounded by $O(n \log n)$. In addition, note that we could find the points $\{\{\gamma_{j,q}\}_{q=1}^{v_j}\}_{j=1}^k$ in $O(n \log n)$ defined in the proof of Lemma 3, by computing $\{u_h^{-1}(u_h(\gamma_j) + \kappa_S)\}$ (as illustrated in Figure 3). Furthermore, the total number of such points, i.e., $\sum_{j=1}^k v_j$, is also k , so this part can also be completed in time $O(n \log n)$.

Next, for any $j \in [k]$, we first find the optimal contract $(\gamma, \beta(\gamma))$, and condition on $\gamma \in [\gamma_j, \gamma_{j+1}]$. After doing so, the principal, among all such contracts, can pick the one that leads to the highest expected utility for her. Therefore, it suffices to focus on the case $\gamma \in [\gamma_j, \gamma_{j+1}]$. Recall that, in this case, the principal's utility is given by

$$\mathcal{U}_p = (1 - \gamma)R_{i_j} - \beta(\gamma)\kappa_I. \quad (12)$$

Note that, by Lemma 3, $\mathcal{U}_p$ is a concave function over $[\gamma_j, \gamma_{j+1}]$. Hence, to find its maximum, we just need to find the point $\frac{d}{d\gamma} \mathcal{U}_p = 0$, where $\beta(\cdot)$ is differentiable, and also check the endpoints and points of nondifferentiable points. Thus, we would need to check for the solutions of

$$\beta'(\gamma) = -\frac{R_{i_j}}{\kappa_I}. \quad (13)$$

Using the notation in the proof of Lemma 3, we have an explicit characterization of $\beta'(\gamma)$ over all intervals $(\gamma_{j,q}, \gamma_{j,q+1})$, and so we can find the potential solution to (13) in time $O(1)$. As a result, the total computation time that we need to check all the intervals $(\gamma_{j,q}, \gamma_{j,q+1})$, for all j and $0 \leq q \leq v_j$, and their endpoints, is $O(n)$. This completes the proof. $\square$

A summary of the above steps is provided in Algorithm 1. It is worth noting that the optimal contract is not always unique.

With the computational guarantees for determining the optimal contracts in hand, we next investigate the comparative statics of the optimal contract, showing how the agent's payment and the probability of inspection vary as the parameters of the model corresponding to the inspection costs change. Note that these results do not require any assumption on the uniqueness of the optimal contract.

THEOREM 2. *Suppose Assumptions 1 and 2 hold. Let (γ^*, β^*) be an optimal contract.*

- (1) *Suppose the principal's cost of inspection κ_I increases, and let (γ'^*, β'^*) denote an optimal contract under this new setting. Then, we have $\gamma'^* \geq \gamma^*$ and $\beta'^* \leq \beta^*$.*
- (2) *Suppose the agent's cost of complying with safety measure κ_S increases, and let (γ''^*, β''^*) denote an optimal contract under this new setting. Then, we have $\gamma''^* \geq \gamma^*$.*

PROOF. Proof of part (1): Notice that changing κ_I does not change the $\{\gamma_j\}_j$ and the $\beta(\cdot)$ function. As a consequence, changing κ_I does not change the actions that (γ^*, β^*) and (γ'^*, β'^*) implement. That said, let us denote the reward of the actions implemented by (γ^*, β^*) and (γ'^*, β'^*) by R and R' , respectively. Also, note that it suffices to show $\beta'^* \leq \beta^*$, and the other result $\gamma'^* \geq \gamma^*$ will be implied using the fact that $\beta(\cdot)$ is a decreasing function.

Suppose κ_I is increased to κ'_I . Since (γ^*, β^*) is the optimal action with κ_I , we have

$$(1 - \gamma^*)R - \beta^* \kappa_I \geq (1 - \gamma'^*)R' - \beta'^* \kappa_I. \quad (14)$$

We prove the desired result by contradiction. Suppose $\beta'^* > \beta^*$. Hence, we have

$$-\beta^* (\kappa'_I - \kappa) > -\beta'^* (\kappa'_I - \kappa). \quad (15)$$

Adding the two sides of (14) and (15) implies

$$(1 - \gamma^*)R - \beta^* \kappa'_I > (1 - \gamma'^*)R' - \beta'^* \kappa'_I. \quad (16)$$

However, this is in contradiction with the assumption that (γ'^*, β'^*) is an optimal contract when the cost of inspection is κ'_I . This completes the proof of this part.

Proof of part (2): Increasing κ_S to κ'_S does not change $\{\gamma_j\}_j$'s and the upper envelope function $u_h(\cdot)$, and therefore, does not change the actions that (γ^*, β^*) and (γ''^*, β''^*) implement. Let us denote the reward of the actions implemented by these two contracts by R and R'' , respectively. On the other hand, changing κ_S changes the $\beta(\cdot)$ function to a new function $\hat{\beta}(\cdot)$. It is straightforward to see that $\hat{\beta}(\cdot)$ is pointwise larger than $\beta(\cdot)$.

Using the optimality of (γ^*, β^*) with κ_S , we have

$$(1 - \gamma^*)R - \beta(\gamma^*) \kappa_I \geq (1 - \gamma''^*)R'' - \beta(\gamma''^*) \kappa_I. \quad (17)$$

The optimality of (γ''^*, β''^*) with κ'_S implies

$$(1 - \gamma''^*)R'' - \hat{\beta}(\gamma''^*) \kappa_I \geq (1 - \gamma^*)R - \hat{\beta}(\gamma^*) \kappa_I. \quad (18)$$

Summing the two sides of (17) and (18) and simplifying it gives us:

$$\hat{\beta}(\gamma^*) - \beta(\gamma^*) \geq \hat{\beta}(\gamma''^*) - \beta(\gamma''^*). \quad (19)$$

We next make the following claim:

CLAIM 1. $\hat{\beta}(\gamma) - \beta(\gamma)$ is a decreasing function of γ .

Notice that this claim, along with (19), gives us the desired result. It remains to show that the claim holds.

Proof of Claim 1: Recall from Lemma 2 that both $\beta(\cdot)$ and $\hat{\beta}(\cdot)$ are decreasing functions of γ . Also, as stated above, for any γ , $\hat{\beta}(\gamma) \geq \beta(\gamma)$. Hence, we could divide $[0, 1]$ to at most three intervals: (1) $[0, \gamma_1^{th})$ where both $\beta(\gamma)$ and $\hat{\beta}(\gamma)$ are positive, (2) $[\gamma_1^{th}, \gamma_2^{th})$ where $\beta(\gamma) = 0$

but $\hat{\beta}(\gamma)$ is positive, and (3) $[\gamma_2^{th}, 1]$ where both $\beta(\gamma)$ and $\hat{\beta}(\gamma)$ are zero. We need to verify that Claim 1 holds over all these three intervals. Note that since $\beta(\gamma)$ and $\hat{\beta}(\gamma)$ are continuous, once we have the claim established in all these three intervals, it also holds over the whole interval $[0, 1]$.

Obviously Claim 1 holds for the third interval $[\gamma_2^{th}, 1]$. It also holds over the second interval, i.e., $[\gamma_1^{th}, \gamma_2^{th}]$, since $\hat{\beta}(\gamma) - \beta(\gamma) = \hat{\beta}(\gamma)$ which is a decreasing function of γ . Hence, it remains to show that our claim holds over the first interval, i.e., when $\beta(\gamma)$ and $\hat{\beta}(\gamma)$ are both positive. In this case $\beta(\gamma)$ is given by

$$\beta(\gamma) = 1 - \frac{u_h^{-1}(u_h(\gamma) - \kappa_S)}{\gamma(1 - \alpha)}. \quad (20)$$

Therefore, we need to show that $g(\gamma, \kappa'_S) - g(\gamma, \kappa_S)$, with

$$g(\gamma, \kappa_S) := \frac{u_h^{-1}(u_h(\gamma) - \kappa_S)}{\gamma(1 - \alpha)}, \quad (21)$$

is an increasing function of γ . Notice that $u_h^{-1}(\cdot)$ is a continuous function which is nondifferentiable at finitely many points. Hence, $g(\gamma, \kappa_S)$, as a function of κ_S , is continuous and nondifferentiable at all but finitely many points. Next, we claim

$$g(\gamma, \kappa'_S) - g(\gamma, \kappa_S) = \int_{\kappa_S}^{\kappa'_S} \frac{\partial}{\partial \kappa} g(\gamma, \kappa) d\kappa, \quad (22)$$

where the integral in (22) is the Lebesgue integral. To establish this result, we can partition the interval $[\kappa_S, \kappa'_S]$ into subintervals where $g(\gamma, \kappa)$ is differentiable with respect to κ over each of them. Subsequently, we can apply the fundamental theorem of calculus to each of these subintervals and then integrate them using the continuity of g to derive (22). Another way to establish (22) is to first recognize that u_h^{-1} is an absolutely continuous function as it is a piecewise linear function. Then we use the generalized version of the fundamental theorem of calculus (see Theorem 7.18 in [Rudin, 1986]).

Now, note that the derivative of g with respect to κ is given by

$$\frac{\partial}{\partial \kappa} g(\gamma, \kappa) = \frac{-1}{\gamma(1 - \alpha)u_h(u_h^{-1}(u_h(\gamma) - \kappa_S))}. \quad (23)$$

Since $u_h(\cdot)$ is increasing function, it is straightforward to see $\frac{\partial}{\partial \kappa} g(\gamma, \kappa)$ is an increasing function of γ . This, along with (22), shows that $g(\gamma, \kappa'_S) - g(\gamma, \kappa_S)$ is an increasing function of γ which completes the proof of Claim 1. $\square$

In essence, the first part of Theorem 2 shows that an increase in the inspection pricing κ_I prompts the principal to reduce the inspection probability. Simultaneously, the principal increases the agent's payment to ensure the agent remains incentivized to observe safety protocols. The second part of the theorem shows that when the agent has to pay higher costs for adhering to safety protocols, the principal increases the agent's payment, thereby motivating continued adherence to safety protocols. In this case, and at first glance, one might intuitively presume that the principal would also increase the inspection probability, based on the same rationale. However, as highlighted in the next example, this is not necessarily the case.

EXAMPLE 1. Consider a setting where agent has $n = 6$ actions with rewards and costs given by

$$[R_i]_{i=1}^6 = [1.5, 3, 4, 6, 7, 9], \quad [c_i]_{i=1}^6 = [1, 1.3, 1.5, 2.5, 3.4, 5.2]. \quad (24)$$

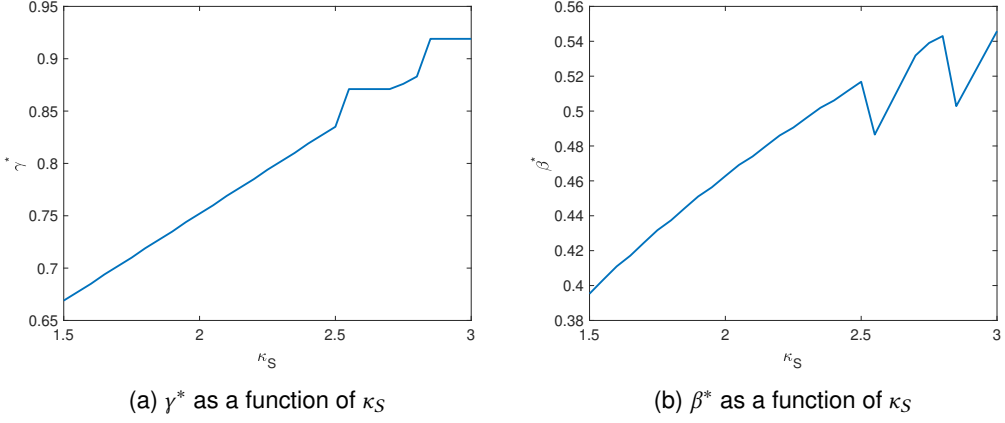


Fig. 5. An example depicting how the optimal contract changes as κ_S increases.

Figure 5 illustrates the optimal share of payment to the agent γ^* and the probability of inspection β^* as we increase κ_S . In particular, as Figure 5a shows γ^* is a (weakly) increasing function of κ_S which is aligned with our result in Theorem 2. On the other hand, as we discussed above and Figure 5b shows, β^* is not necessarily a monotone function of the cost κ_S . The underlying cause of this phenomenon is the interplay of two opposing forces. On one hand, by increasing the cost of safety for the agent, i.e., κ_S , the beta function $\beta(\cdot)$ increases pointwise. Put another way, had the payment ratio γ remained unchanged, the probability of inspection would have gone up. However, as we established earlier, the optimal γ^* at equilibrium may also increase. In this context, an increased payment implies that we reduce the probability of inspection. The combined impact of these two forces dictates whether the optimal probability of inspection β^* ascends or descends as a function of κ_S .

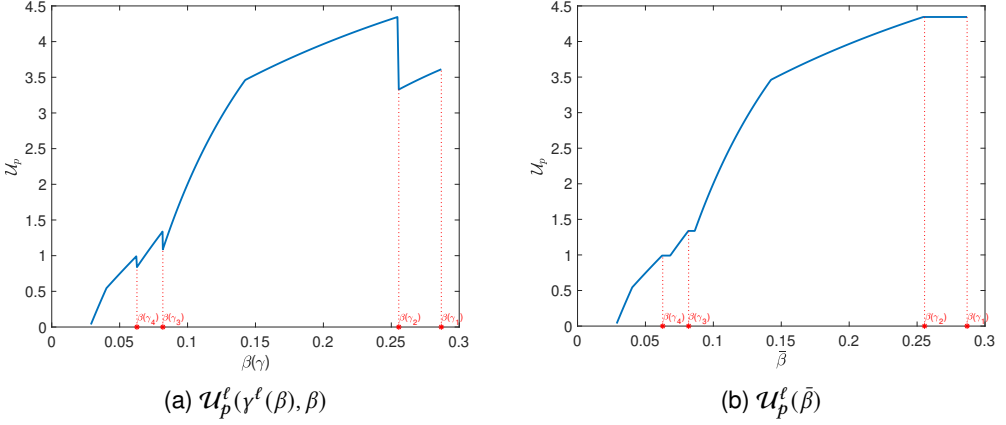
4 The Multi-Contract Setting

We turn to the multi-contract setting. Although we understand how to determine the optimal contract for each agent, we must ensure that the cumulative probability of inspection does not exceed the inspection budget. For every agent $\ell \in [m]$, let $\beta_{\min}^\ell$ represent the lowest possible probability of inspection across all contracts that implement one of the actions for agent j . As established by Lemma 2, the probability of inspection diminishes as the agent's payment increases. Thus, it reaches its minimum when the agent receives the maximum payment, which is the total reward. Consequently, we deduce that $\beta_{\min}^\ell = \beta^\ell(1)$, where $\beta^\ell(\cdot)$ is the $\beta(\cdot)$ function introduced in Proposition 1 in relation to agent ℓ .² The subsequent assumption ensures the presence of at least one feasible set of contracts.

ASSUMPTION 3. We have $\sum_{\ell=1}^m \beta_{\min}^\ell \leq B$.

Using Theorem 1, we can find optimal contracts for different agents separately and do so in time $O(mn \log n)$. However, the total probability of inspection could potentially exceed B . In such cases, compromises would be necessary, meaning we would have to consider suboptimal contracts that can be implemented with a lower probability of inspection than

²We reuse the major notation from the single-agent setting by using superscripts to distinguish among the different agents. In particular, recalling Proposition 1, and for any agent $\ell \in [m]$, we have $0 < \gamma_1^\ell < \gamma_2^\ell < \dots < \gamma_{k^\ell}^\ell < \gamma_{k^\ell+1}^\ell = 1$, where for any $j \in [k^\ell]$, the action $(a_{ij}^\ell, 1)$ is implementable for agent ℓ with $\gamma \in [\gamma_j^\ell, \gamma_{j+1}^\ell]$.

Fig. 6. Principal's utility for agent's ℓ action

the optimal ones. Consequently, the following natural question arises: *For any agent $\ell \in [m]$, and given some $\bar{\beta}^\ell$, which contract maximizes the principal's utility $\mathcal{U}_p^\ell$ among all those whose probability of inspection is at most $\bar{\beta}^\ell$?*

Notice that, with $\gamma \in [\gamma_j^\ell, \gamma_{j+1}^\ell]$, the utility of principal from agent ℓ 's action is given by

$$\mathcal{U}_p^\ell(\gamma, \beta^\ell(\gamma)) = (1 - \gamma)R_{i_j^\ell}^\ell - \beta^\ell(\gamma)\kappa_I^\ell. \quad (25)$$

For the sake of our analysis, we find it more convenient to interpret the principal's utility as a function of β . Let us denote $\beta(\gamma_j^\ell)$ by β_j^ℓ . Notice that, since $\beta^\ell(\cdot)$ is a decreasing function, we have

$$\beta_1^\ell > \dots > \beta_{k^\ell}^\ell > \beta_{k^\ell+1}^\ell = \beta_{\min}^\ell. \quad (26)$$

We also denote the inverse of $\beta^\ell(\cdot)$ by $\gamma^\ell(\cdot)$. Using Lemma 2 and 3, it is straightforward to verify that $\gamma^\ell(\cdot)$ is also a decreasing function and it is convex over each interval $[\beta_{j+1}^\ell, \beta_j^\ell]$.

Next, we can rewrite the principal's utility from agent ℓ 's action given in (25) as a function of β . In particular, for any $\beta \in [\beta_{j+1}^\ell, \beta_j^\ell]$, we have

$$\mathcal{U}_p^\ell(\gamma^\ell(\beta), \beta) = (1 - \gamma^\ell(\beta))R_{i_j^\ell}^\ell - \beta\kappa_I^\ell. \quad (27)$$

This is a concave function over each interval $[\beta_{j+1}^\ell, \beta_j^\ell]$. Figure 6a depicts this function for the example provided in Remark 1. The dashed lines highlight the intervals $[\beta_{j+1}^\ell, \beta_j^\ell]$'s.

Now, going back to our question above, with a slight abuse of notation, we denote the maximum utility that the principal can obtain from agent ℓ given the condition $\beta \leq \bar{\beta}$ by $\mathcal{U}_p^\ell(\bar{\beta})$ which is given by

$$\mathcal{U}_p^\ell(\bar{\beta}) = \max_{\beta \leq \bar{\beta}} \mathcal{U}_p^\ell(\gamma^\ell(\beta), \beta). \quad (28)$$

It is straightforward to see that this function is (weakly) increasing. Moreover, since $\mathcal{U}_p^\ell(\gamma^\ell(\beta), \beta)$ is concave over each interval $[\beta_{j+1}^\ell, \beta_j^\ell]$, the function $\mathcal{U}_p^\ell(\bar{\beta})$ consists of segments that are either constant or both concave and increasing. Figure 6b illustrates this function $\mathcal{U}_p^\ell(\bar{\beta})$ for the same example of Figure 6a.

Now, the principal's problem can be formulated as

$$\max_{(\tilde{\beta}^\ell)_{\ell=1}^m} \sum_{\ell=1}^m \mathcal{U}_p^\ell(\tilde{\beta}^\ell) \text{ such that } \sum_{\ell=1}^m \tilde{\beta}^\ell \leq B. \quad (29)$$

The optimization problem (29) bears a resemblance to the multiple-choice knapsack problem (MCKP), a variant of the classic knapsack problem. In the MCKP, items are categorized into classes, and the goal is to maximize the cumulative value of the chosen items without surpassing the knapsack's capacity and under the constraint of selecting at most one item from each class. To make the connection to our setting clearer, consider discretizing each function $\mathcal{U}_p^\ell(\cdot)$ and grouping all the samples into a single class. For a given class $\ell \in [m]$, each item takes the form $(\beta, \mathcal{U}_p^\ell(\beta))$, where the probability of inspection β is seen as the weight and $\mathcal{U}_p^\ell(\beta)$ is interpreted as this item's value. The knapsack's capacity in our scenario is set to B , representing the total inspection budget. The constraint of selecting one item from each class translates to the constraint that we can inspect each agent at most once.

A comprehensive survey of various methods to address the MCKP can be found in [Kellerer et al., 2004]. We choose to use a dynamic programming approach which is inspired by the algorithm presented in [Dudziński and Walukiewicz, 1987] for the MCKP. This yields the following complexity result.

THEOREM 3. *Suppose Assumptions 1-3 hold. Then, for any $\epsilon > 0$, an ϵ -approximate solution to the principal's problem (29) can be found in time $O\left(mn + \frac{m^3 B}{\epsilon^2} \log n\right)$.*

PROOF. First recall that, similar to Algorithm 1, we can characterize the function $\mathcal{U}_p^\ell(\cdot)$ in time $O(n \log n)$ for any ℓ , and hence, in time $O(mn \log n)$ for all $\ell \in [m]$. Consequently, we can compute $\mathcal{U}_p^\ell(\beta)$ at any given β in time $O(\log n)$.

To ensure each agent receives the minimum inspection, we rewrite the optimization problem (29) as

$$\max_{(x^\ell)_{\ell=1}^m} \sum_{\ell=1}^m \mathcal{U}_p^\ell(\beta_{\min}^\ell + x^\ell) \text{ such that } \sum_{\ell=1}^m x^\ell \leq \tilde{B} := B - \sum_{\ell=1}^m \beta_{\min}^\ell. \quad (30)$$

We next discretize the inspection levels with stepsize $\delta > 0$. Let us denote the grid corresponding to agent ℓ by $\mathcal{G}^\ell$. For any $l \in [m]$ and nonnegative $j \leq \frac{\tilde{B}}{\delta}$, let $V(l, j)$ be the solution to the following maximization problem:

$$V(l, j) := \max_{(x^\ell \in \mathcal{G}^\ell)_{\ell=1}^l} \sum_{\ell=1}^l \mathcal{U}_p^\ell(\beta_{\min}^\ell + x^\ell) + \sum_{\ell=l+1}^m \mathcal{U}_p^\ell(\beta_{\min}^\ell) \text{ such that } \sum_{\ell=1}^l x^\ell \leq j\delta. \quad (31)$$

In other words, $V(l, j)$ represents the highest utility the principal can achieve when searching over the grid, provided they only consider the first l agents for any inspections beyond the minimum and allocate only $j\delta$ from their additional inspection budget.

Note that $V(l, j)$ admits the following recursive characterization:

$$V(l, j) = \max_{\eta=0, \dots, \min\{j, 1/\delta\}} \left(V(l-1, j-\eta) + \mathcal{U}_p^l(\beta_{\min}^l + \eta\delta) - \mathcal{U}_p^l(\beta_{\min}^l) \right). \quad (32)$$

Using (32), $V(l, j)$ can be computed in time $O(\log n / \delta)$. As a result, we can compute $V(m, \lfloor \tilde{B} / \delta \rfloor)$ in time $O(m\tilde{B} / \delta^2 \log n)$ (and by going back recursively, we find the corresponding optimal probability of inspection for each agent). Finally, we bound the error of such a discretization.

LEMMA 4. Let OPT denote the solution to (30). Then, the solution found over the grid $\prod_{\ell=1}^m \mathcal{G}^\ell$ with stepsize δ is lower bounded by

$$OPT - \delta \left(\sum_{\ell=1}^m \left[\frac{(R_n^\ell)^2}{\kappa_S^\ell} - \kappa_I^\ell \right] \right). \quad (33)$$

We defer the proof of the lemma to the appendix. Given this result, by setting

$$\delta = \epsilon \mathcal{U}_p^1(B - \sum_{\ell=2}^m \beta_{\min}^\ell) \left(m \max_{\ell} \frac{(R_n^\ell)^2}{\kappa_S^\ell} \right)^{-1}, \quad (34)$$

in which $\mathcal{U}_p^1(B - \sum_{\ell=2}^m \beta_{\min}^\ell)$ serves as a lower bound for the OPT , we obtain the desired approximation. $\square$

Having obtained the approximately optimal solution, a natural question arises regarding its implementation: Given a specific allocation of the inspection budgets $(\tilde{\beta}^\ell)_{\ell=1}^m$, how should the B inspectors be (randomly) allocated among the m agents to ensure that agent ℓ is inspected with probability $\tilde{\beta}^\ell$? When $B = 1$, the solution is straightforward: one can generate a uniform random variable U over $[0, 1]$ and inspect agent ℓ if U falls within the interval $[\sum_{i=1}^{\ell-1} \tilde{\beta}^i, \sum_{i=1}^{\ell} \tilde{\beta}^i]$. However, when B is greater than one, the situation becomes more complex because we must ensure that each agent is inspected by at most one inspector. We can formulate this problem within the framework of the well-known Birkhoff-von-Neumann algorithm [Birkhoff, 1946] by constructing a matrix where each entry represents the probability that a specific agent is inspected by a particular inspector. This algorithm presents a method to decompose an $m \times m$ bistochastic matrix into a convex combination of permutation matrices, with a time complexity of $O(m^2)$. However, in our setting, there are no constraints on the joint distribution of inspectors and agents, apart from the provided marginals. This allows us to derive an intuitive and simpler algorithm that runs in $O(m)$ time complexity.

LEMMA 5. For any given vector of $(\tilde{\beta}^\ell)_{\ell=1}^m$, there exists a random algorithm to assign inspectors to agents in time $O(m)$ such that each agent $\ell \in [m]$ is inspected with probability $\tilde{\beta}^\ell$.

Proof sketch: Consider the first inspector who wants to allocate their one unit of inspection across agents. They start with agent one, dedicating a $\tilde{\beta}^1$ fraction of their inspection budget, and then proceed to agent two, continuing in this manner until their budget is over. This procedure, however, implies that the final agent they inspected (let us call it agent ℓ) might be inspected at a probability lower than the target, $\tilde{\beta}^\ell$, due to the depletion of their inspection unit. Consequently, the second inspector should start from this agent, taking care of the residual inspection probability, and then advance to subsequent agents.

It is critical to ensure that agent ℓ undergoes inspection by no more than a single inspector. We design a joint inspection plan for the initial two inspectors ensuring that, in each instance, at most one of them inspects agent ℓ . In essence, the inspection schedule of the second inspector is contingent on the actions of the first. The second inspector is allowed to inspect agent ℓ only when the first inspector had inspected one of the initial $\ell - 1$ agents. The scheduling for subsequent inspectors is designed similarly, ensuring that in cases where two consecutive inspectors are in charge of inspecting one agent, they do not conduct the inspection simultaneously. The details of this method are provided in the appendix for the sake of completeness. $\square$

4.1 An Illustrative Example

In general, the structure of the optimal allocation of the inspection budget across agents could be complex or even counter-intuitive. For example, even when we have a set of homogeneous agents, the optimal inspection allocation does not necessarily involve inspecting all agents equally. To illustrate, consider Figure 6b. Imagine we have two identical agents whose corresponding principal utilities are depicted in this figure. Specifically, the utility of the principal from one agent when inspected with probability $\beta_{\min} + 0.1$ is above 3, which is more than twice the utility from one agent when inspected with probability $\beta_{\min} + 0.05$, a value below 1.3. This suggests that, given two homogeneous agents with these characteristics, inspecting both equally could be suboptimal compared to inspecting one at the minimum level and allocating all the extra inspection budget to the other.

We next provide a case study that illuminates the phenomenon where the principal might opt to forgo inspecting a subset of agents, especially when there is a disparity in inspection costs. Imagine a scenario in which agents are identical in all respects except for their inspection costs. Specifically, assume that the inspection costs for agents take on one of two distinct values: high and low. Half of the agents have high inspection costs, while the other half have low costs. In this case, we show the following result:

PROPOSITION 2. *Given the scenario described above and assuming $\beta_{\min}^{\ell} = 0$ for all $\ell \in [m]$, there exists a threshold M such that for every $m \geq M$, the principal does not inspect half of the agents associated with a higher inspection cost. Specifically, in any optimal solution to the principal's optimization problem (29), these agents' allocated level of inspection is zero.*

The proof can be found in the appendix. This outcome emphasizes that as the number of agents increases, the principal may choose to not monitor a substantial subset of agents that entail high inspection costs, directing its entire inspection budget towards those that are less costly to inspect. Conversely, this suggests that agents who manage to elevate the inspection costs for the principal might escape the inspection and, as a result, secure a higher payment.

5 Conclusion

We study the role of safety and regulatory inspections in contract design problems. In particular, we consider a principal who can use random and costly inspections, in addition to payments, to incentivize agents to take safe actions. For the single-agent setting, we provide an efficient algorithm to find the optimal linear contract, and also establish how the payment fraction and inspection probability vary as the costs of inspection or adherence to safety protocols increase. We extend our results to the multi-agent setting, where we draw connections with the Knapsack problem to find an approximately optimal set of contracts in that case. Our case study on the structure of the solution illustrates how agents who are most costly to inspect may escape monitoring.

We believe our framework can be used and extended to study further problems surrounding regulatory actions in contract design. In particular, one interesting future direction would be to study the dynamic setting where the principal interacts with the agents across multiple rounds and must decide how often to inspect different agents.

Acknowledgments

The authors thank Alex Teytelboym, Ali Makhdoumi, and Azarakhsh Malekian for insightful discussions and comments. Alireza Fallah acknowledges support by the National Science Foundation under grant number DMS-1928930 and by the Alfred P. Sloan Foundation under

grant G-2021-16778, during his residence at the Simons Laufer Mathematical Sciences Institute in Berkeley, California, during the Fall 2023 semester. Michael Jordan acknowledges support from the Mathematical Data Science program of the Office of Naval Research under grant number N00014-21-1-2840 and the European Research Council Synergy Program.

References

- Saeed Alaei, Alexandre Belloni, Ali Makhdoumi, and Azarakhsh Malekian. 2020. Optimal Auction Design with Deferred Inspection and Reward. *Available at SSRN 3700525* (2020).
- Moshe Babaioff, Michal Feldman, and Noam Nisan. 2006. Combinatorial agency. In *Proceedings of the 7th ACM Conference on Electronic Commerce*. 18–28.
- Ian Ball and Deniz Kattwinkel. 2019. Probabilistic verification in mechanism design. In *Proceedings of the 2019 ACM Conference on Economics and Computation*. 389–390.
- Ian Ball and Jan Knoepfle. 2023. Should the Timing of Inspections be Predictable? *arXiv preprint arXiv:2304.01385* (2023).
- Andrei Barbos. 2022. Optimal contracts with random monitoring. *International Journal of Game Theory* 51, 1 (2022), 119–154.
- Elchanan Ben-Porath, Eddie Dekel, and Barton L Lipman. 2014. Optimal allocation with costly verification. *American Economic Review* 104, 12 (2014), 3779–3813.
- Dimitri P Bertsekas. 1997. Nonlinear programming. *Journal of the Operational Research Society* 48, 3 (1997), 334–334.
- Garrett Birkhoff. 1946. Tres observaciones sobre el algebra lineal. *Univ. Nac. Tucuman, Ser. A* 5 (1946), 147–154.
- Patrick Bolton and Mathias Dewatripont. 2004. *Contract Theory*. MIT press.
- Ioannis Caragiannis, Edith Elkind, Mario Szegedy, and Lan Yu. 2012. Mechanism design: from partial to probabilistic verification. In *Proceedings of the 13th acm conference on electronic commerce*. 266–283.
- Gabriel Carroll. 2015. Robustness and linear contracts. *American Economic Review* 105, 2 (2015), 536–563.
- Mingliu Chen, Peng Sun, and Yongbo Xiao. 2020. Optimal monitoring schedule in dynamic contracts. *Operations Research* 68, 5 (2020), 1285–1314.
- Chongwoo Choe and Iain Fraser. 1999. Compliance monitoring and agri-environmental policy. *Journal of agricultural economics* 50, 3 (1999), 468–487.
- Krzysztof Dudziński and Stanisław Walukiewicz. 1987. Exact methods for the knapsack problem and its generalizations. *European Journal of Operational Research* 28, 1 (1987), 3–21.
- Paul Dütting, Tomer Ezra, Michal Feldman, and Thomas Kesselheim. 2022. Combinatorial contracts. In *2021 IEEE 62nd Annual Symposium on Foundations of Computer Science (FOCS)*. IEEE, 815–826.
- Paul Dütting, Tim Roughgarden, and Inbal Talgam-Cohen. 2019. Simple versus optimal contracts. In *Proceedings of the 2019 ACM Conference on Economics and Computation*. 369–387.
- Paul Dütting, Tim Roughgarden, and Inbal Talgam-Cohen. 2021. The complexity of contracts. *SIAM J. Comput.* 50, 1 (2021), 211–254.
- Albin Erlanson and Andreas Kleiner. 2020. Costly verification in collective decisions. *Theoretical Economics* 15, 3 (2020), 923–954.
- Diodato Ferraioli and Carmine Ventre. 2018. Probabilistic verification for obviously strategyproof mechanisms. *arXiv preprint arXiv:1804.10512* (2018).
- Paul J Ferraro. 2008. Asymmetric information and contract design for payments for environmental services. *Ecological Economics* 65, 4 (2008), 810–821.
- Douglas Gale and Martin Hellwig. 1985. Incentive-compatible debt contracts: The one-period problem. *The Review of Economic Studies* 52, 4 (1985), 647–663.
- Ronald L. Graham. 1972. An efficient algorithm for determining the convex hull of a finite planar set. *Info. Proc. Lett.* 1 (1972), 132–133.
- Jerry R Green and Jean-Jacques Laffont. 1986. Partially verifiable information and mechanism design. *The Review of Economic Studies* 53, 3 (1986), 447–456.
- Sanford J Grossman and Oliver D Hart. 1992. An analysis of the principal-agent problem. In *Foundations of Insurance Economics: Readings in Economics and Finance*. Springer, 302–340.
- Marina Halac and Pierre Yared. 2020. Commitment versus flexibility with costly verification. *Journal of Political Economy* 128, 12 (2020), 4523–4573.
- Winston Harrington. 1988. Enforcement leverage when penalties are restricted. *Journal of Public Economics* 37, 1 (1988), 29–53.
- Peter-Jürgen Jost. 1991. Monitoring in principal-agent relationships. *Journal of Institutional and Theoretical Economics (JITE)/Zeitschrift für die gesamte Staatswissenschaft* (1991), 517–538.

- Peter-Jürgen Jost. 1996. On the role of commitment in a principal–agent relationship with an informed principal. *Journal of Economic Theory* 68, 2 (1996), 510–530.
- Hans Kellerer, Ulrich Pferschy, David Pisinger, Hans Kellerer, Ulrich Pferschy, and David Pisinger. 2004. The multiple-choice knapsack problem. *Knapsack Problems* (2004), 317–347.
- Yunan Li. 2020. Mechanism design with costly verification and limited punishments. *Journal of Economic Theory* 186 (2020), 105000.
- Tymofiy Mylovanov and Andriy Zapechelnyuk. 2017. Optimal allocation with ex post verification and limited penalties. *American Economic Review* 107, 9 (2017), 2666–2694.
- Dmitry Orlov. 2022. Frequent monitoring in dynamic contracts. *Journal of Economic Theory* 206 (2022), 105550.
- Walter Rudin. 1986. *Real and Complex Analysis* (3rd ed.). McGraw-Hill, New York.
- Roland Strausz. 1997. Delegation of monitoring in a principal-agent relationship. *The Review of Economic Studies* 64, 3 (1997), 337–357.
- Robert M Townsend. 1979. Optimal contracts and competitive markets with costly state verification. *Journal of Economic Theory* 21, 2 (1979), 265–293.
- Felipe Varas, Iván Marinovic, and Andrzej Skrzypacz. 2020. Random inspections and periodic reviews: Optimal dynamic monitoring. *The Review of Economic Studies* 87, 6 (2020), 2893–2937.
- Banghua Zhu, Stephen Bates, Zhuoran Yang, Yixin Wang, Jiantao Jiao, and Michael I. Jordan. 2023. The Sample Complexity of Online Contract Design. In *Proceedings of the 24th ACM Conference on Economics and Computation* (London, United Kingdom) (EC '23). Association for Computing Machinery, New York, NY, USA, 1188.

A Deferred Proofs

A.1 Proof of Lemma 1

Suppose some safe action $(a_i, 1)$ is implementable. Then, by IC, we should have

$$\mathcal{U}_a(a_i, 1) \geq \mathcal{U}_a(a_i, 0), \quad (\text{A1})$$

which implies

$$\gamma R_i - c_i - \kappa_S \geq (1 - \alpha)\gamma R_i - c_i. \quad (\text{A2})$$

As a consequence, we should have

$$\alpha\gamma R_i \geq \kappa_S. \quad (\text{A3})$$

Using $1 \geq \gamma$ and $R_n \geq R_i$, implies that $\alpha R_n \geq \kappa_S$ which contradicts the assumption given in the lemma's statement.

A.2 Proof of Lemma 2

As depicted in Figure 2, we can cast $\beta(\gamma)$ as:

$$\beta(\gamma) = \max \left\{ 1 - \frac{u_h^{-1}(u_h(\gamma) - \kappa_S)}{\gamma(1 - \alpha)}, 0 \right\}. \quad (\text{A4})$$

Hence, it suffices to show that

$$f(\gamma) := \frac{u_h^{-1}(u_h(\gamma) - \kappa_S)}{\gamma(1 - \alpha)} \quad (\text{A5})$$

is strictly increasing in γ . Note that $u_h(\cdot)$ is a piecewise linear function, and therefore, it is differentiable on all but finitely many points. As a result, and due to the strict monotonicity of $u_h(\cdot)$, the function $f(\cdot)$ is likewise differentiable, except at a finite number of points. Consequently, there exists a sequence $0 = \rho_1 < \rho_2 < \dots < \rho_v = 1$ such that f is differentiable within each interval (ρ_j, ρ_{j+1}) . Furthermore, since $u_h(\cdot)$ and its inverse are both continuous, $f(\cdot)$ is also continuous. We claim that it suffices to show that the derivative of $f(\cdot)$ is positive wherever it is differentiable. By proving this, we establish that f is increasing on each interval (ρ_j, ρ_{j+1}) , which, combined with the continuity of f , concludes our proof.

To show the aforementioned claim holds, note that the derivative of f is given by:

$$f'(\gamma) = \frac{\frac{u'_h(\gamma)}{u'_h(u_h^{-1}(u_h(\gamma) - \kappa_S))} \gamma - u_h^{-1}(u_h(\gamma) - \kappa_S)}{\gamma^2(1 - \alpha)}. \quad (\text{A6})$$

Let us denote $u_h^{-1}(u_h(\gamma) - \kappa_S)$ by $\tilde{\gamma}$. To show the derivative (A6) is positive, we need to show that $\gamma u'_h(\gamma) > \tilde{\gamma} u'_h(\tilde{\gamma})$. First, notice that, since $u_h(\cdot)$ is increasing, $\tilde{\gamma} < \gamma$. Second, notice that $u_h(\cdot)$ is convex as it is the maximum of a collection of linear (and hence convex) functions. Consequently, $u'_h(\cdot)$ is increasing, and thus, $u'_h(\gamma) > u'_h(\tilde{\gamma})$ as well, which completes the proof.

A.3 Proof of Lemma 3

Recall the definition of function f from the proof of Lemma 2. It suffices to show $f(\cdot)$ is strictly concave over the interval $[\gamma_j, \gamma_{j+1}]$. To do so, we first show $f(\cdot)$ is strictly concave over the interval $(\gamma_{j,q}, \gamma_{j,q+1})$ for any q .

Let $\gamma \in (\gamma_{j,q}, \gamma_{j,q+1})$ and suppose $\tilde{\gamma}$ falls within the interval $[\gamma_{j,q}, \gamma_{j,q+1}]$ where $j_q \leq j$. Recall the derivative of f in (A6) is given by

$$f'(\gamma) = \frac{\frac{u'_h(\gamma)}{u'_h(\tilde{\gamma})} \gamma - \tilde{\gamma}}{\gamma^2(1-\alpha)}. \quad (\text{A7})$$

Note that $u'_h(\gamma) = R_{ij}$ and $u'_h(\tilde{\gamma}) = R_{ijq}$. Therefore, we can rewrite the derivative of f at γ as:

$$f'(\gamma) = \frac{R_{ij}\gamma - R_{ijq}\tilde{\gamma}}{R_{ijq}\gamma^2(1-\alpha)}. \quad (\text{A8})$$

Recall that $u_h(\tilde{\gamma}) = R_{ijq}\tilde{\gamma} - c_{ijq}$ is equal to $u_h(\gamma) - \kappa_S$ where $u_h(\gamma) = R_{ij}\gamma - c_{ij}$. As a result, we can simplify the numerator of (A8) to $c_{ij} - c_{ijq} + \kappa_S$. Thus, we have

$$f'(\gamma) = \frac{c_{ij} - c_{ijq} + \kappa_S}{R_{ijq}\gamma^2(1-\alpha)}. \quad (\text{A9})$$

Therefore, $f'(\gamma)$ is decreasing over $(\gamma_{j,q}, \gamma_{j,q+1})$, and hence, $f'(\cdot)$ is concave over this interval. Moreover, as we go from one interval to another, i.e., as q increases, R_{ijq} and c_{ijq} both increase, meaning that $f'(\gamma)$ further decreases. This, along with the fact that minimum of concave functions is also concave, implies that $f(\cdot)$ is concave over the whole interval $[\gamma_j, \gamma_{j+1}]$. This completes our proof.

A.4 Proof of Lemma 4 in Theorem 3

It is sufficient to prove that for any ℓ , $\mathcal{U}_p^\ell(\cdot)$ is Lipschitz with a parameter bounded by

$$\frac{(R_n^\ell)^2}{\kappa_S^\ell} - \kappa_I^\ell.$$

By taking the optimal solution and rounding agent ℓ 's inspection level down to the nearest element of the grid $\mathcal{G}^\ell$, the discretization error will be limited to $\delta \left\lceil \frac{(R_n^\ell)^2}{\kappa_S^\ell} - \kappa_I^\ell \right\rceil$ given that the grid has a resolution of δ . Note that by rounding down, we ensure the inspection budget is not exceeded.

Further, according to (28), $\mathcal{U}_p^\ell(\beta)$ is either equal to $\mathcal{U}_p^\ell(\gamma^\ell(\beta), \beta)$ or remains constant. Hence, to determine the upper bound on the Lipschitz parameter of $\mathcal{U}_p^\ell(\cdot)$, it is enough to find the upper bound for the derivative of $\mathcal{U}_p^\ell(\gamma^\ell(\beta), \beta)$ as a function of β . While it might not be a continuously differentiable function, it is piecewise differentiable. Given its continuity, bounding its derivative across each segment suffices for our objective. Recall from (27)

$$\mathcal{U}_p^\ell(\gamma^\ell(\beta), \beta) = (1 - \gamma^\ell(\beta))R_{ij}^\ell - \beta\kappa_I^\ell,$$

where $\gamma^\ell(\beta)$ is a decreasing function of β . Hence, when differentiable, we have

$$\frac{d}{d\beta} \mathcal{U}_p^\ell(\gamma^\ell(\beta), \beta) \leq \left| \frac{d}{d\beta} \gamma^\ell(\beta) \right| R_n^\ell - \kappa_I^\ell, \quad (\text{A10})$$

where we used the fact that $R_n^\ell \geq R_{ij}^\ell$. Next, using the inverse function theorem, we have

$$\left| \frac{d}{d\beta} \gamma^\ell(\beta) \right| \leq \left| \frac{1}{f'^\ell(\gamma)} \right|_{\text{at } \gamma^\ell(\beta)}, \quad (\text{A11})$$

where $f^\ell(\cdot)$ is defined for agent ℓ and similar to (A5) in the proof of Lemma 2 for the single-agent case. Note that, by (A9), we have

$$f'^\ell(\gamma) \geq \frac{\kappa_S^\ell}{R_n^\ell \gamma^2 (1 - \alpha)} \geq \frac{\kappa_S^\ell}{R_n^\ell},$$

which, along with (A10) completes the proof of the lemma.

A.5 Proof of Lemma 5

For any $b \in [B]$, define ℓ_b as the smallest integer ℓ such that $\sum_{j=1}^\ell \bar{\beta}^j \geq b$ (and let $\ell_0 = 1$ for convenience).

Inspectors' assignments are decided sequentially, progressing from inspector 1 to B . Specifically, inspector b is assigned to agent $w_b \in \{\ell_{b-1}, \ell_{b-1} + 1, \dots, \ell_b\}$. Starting with the first inspector, they inspect agent $\ell < \ell_1$ with probability $\bar{\beta}^\ell$ and agent ℓ_1 with probability $\zeta_1 := 1 - \sum_{i=1}^{\ell_1-1} \bar{\beta}^i$. This assignment can be carried out using the uniform random variable generator as stated earlier.

Assuming we have determined assignments for inspectors $1, \dots, b-1$, we next decide the agent for inspector b . A key challenge arises because agent ℓ_{b-1} is inspected with probability $\zeta_{b-1} := b-1 - \sum_{i=1}^{\ell_{b-1}-1} \bar{\beta}^i$, which could be less than its targeted inspection probability $\bar{\beta}^{\ell_{b-1}}$. Therefore, inspector b should inspect agent ℓ_{b-1} with the remaining probability $\bar{\beta}^{\ell_{b-1}} - \zeta_{b-1}$, but this should only occur when inspector $b-1$ has opted to inspect other agents. To do so, we use the following procedure:

- i) If $w_{b-1} = \ell_{b-1}$, then inspector b inspects one of the agents in the set $\{\ell_{b-1} + 1, \dots, \ell_b\}$. In other words, in this scenario, inspector b 's single unit of inspection is distributed among these agents. In particular, for agent $\ell < \ell_b$, their probability of inspection is proportional to $\bar{\beta}^\ell$, and for agent ℓ_b , it is proportional to ζ_b . More formally, agent $\ell \in \{\ell_{b-1} + 1, \dots, \ell_b - 1\}$ is inspected with probability

$$\frac{\bar{\beta}^\ell}{1 - \bar{\beta}^{\ell_{b-1}} + \zeta_{b-1}},$$

and agent ℓ_b is inspected with probability

$$\frac{\zeta_b}{1 - \bar{\beta}^{\ell_{b-1}} + \zeta_{b-1}}.$$

- ii) If $w_{b-1} \neq \ell_{b-1}$, then agent ℓ_{b-1} is inspected by probability

$$\frac{\bar{\beta}^{\ell_{b-1}} - \zeta_{b-1}}{1 - \zeta_{b-1}}.$$

The remaining probability of inspection in this case is divided between agents $\{\ell_{b-1} + 1, \dots, \ell_b\}$ in a proportional way similar in part (i).

The above procedure ensures that agent ℓ_{b-1} is not inspected by both inspectors $b-1$ and b . Also, it is straightforward to verify that each agent ℓ 's probability of inspection matches the given desired level $\bar{\beta}^{\ell_{b-1}}$. This completes the proof.

A.6 Proof of Proposition 2

For simplicity of notation, assume that m is even. All results proceed similarly if m is odd. We denote the high and low inspection costs by κ_I^H and κ_I^L , respectively. Without loss of

generality, let us assume the inspection cost corresponding to the first $m/2$ agents is κ_I^L and the inspection cost corresponding to the second half is κ_I^H .

Since agents differ only in their inspection costs, we deduce that

$$\mathcal{U}_p^1(\gamma^1(\beta), \beta) = \dots = \mathcal{U}_p^{m/2}(\gamma^{m/2}(\beta), \beta) > \mathcal{U}_p^{m/2+1}(\gamma^{m/2+1}(\beta), \beta) = \dots = \mathcal{U}_p^m(\gamma^m(\beta), \beta),$$

for any $\beta > 0$. This leads directly to

$$\mathcal{U}_p^1(\beta) = \dots = \mathcal{U}_p^{m/2}(\beta) > \mathcal{U}_p^{m/2+1}(\beta) = \dots = \mathcal{U}_p^m(\beta) \quad (\text{A12})$$

for any $\beta > 0$.

Let $(\bar{\beta}^1, \dots, \bar{\beta}^m)$ be a solution to the optimization problem (29). Using (A12), we infer $\bar{\beta}^i \geq \bar{\beta}^j$ for any $i \in [m/2]$ and $j \in \{m/2 + 1, \dots, m\}$. The inequality becomes strict when the values are positive. To see this, note that, otherwise, the principal could swap $\bar{\beta}^i$ with $\bar{\beta}^j$ to increase their utility.

Without loss of generality, let us assume that $\bar{\beta}^{m/2}$ is the minimum inspection level among the first half of the agents, i.e., $\bar{\beta}^{m/2} = \min \bar{\beta}^1, \dots, \bar{\beta}^{m/2}$, and $\bar{\beta}^{m/2+1}$ is the maximum inspection level among the second half, i.e., $\bar{\beta}^{m/2+1} = \max \bar{\beta}^{m/2+1}, \dots, \bar{\beta}^m$. As previously discussed, $\bar{\beta}^{m/2} \geq \bar{\beta}^{m/2+1}$, and the inequality is strict if the value on the left-hand side is positive.

Now, if $\bar{\beta}^{m/2+1} = 0$, then we are done. Otherwise, we have

$$\frac{2B}{m} \geq \bar{\beta}^{m/2} > \bar{\beta}^{m/2+1} > 0.$$

Given that $\mathcal{U}_p^\ell(\cdot)$ is piecewise differentiable, for sufficiently large m , both $\bar{\beta}^{m/2}$ and $\bar{\beta}^{m/2+1}$ will fall in the first differentiable segments of $\mathcal{U}_p^{m/2}(\cdot)$ and $\mathcal{U}_p^{m/2+1}(\cdot)$ respectively. As a result, and given that agents are similar aside from their inspection costs, for $l \in \{m/2, m/2 + 1\}$ we have

$$\frac{d}{d\beta} \mathcal{U}_p^l(\bar{\beta}^l) = -\frac{d}{d\beta} \gamma(\beta)R - \kappa_I^\ell, \quad (\text{A13})$$

for some R and $\gamma(\cdot)$.

On the other hand, note that we have

$$\frac{d}{d\beta} \mathcal{U}_p^{m/2}(\bar{\beta}^{m/2}) = \frac{d}{d\beta} \mathcal{U}_p^{m/2+1}(\bar{\beta}^{m/2+1}), \quad (\text{A14})$$

because if this were not the case, a slight increase in the inspection level of the one with the higher derivative, along with the decrease by the same amount in the one with the lower derivative, would boost the total principal utility. (The Karush–Kuhn–Tucker (KKT) conditions can also be used to arrive at this result, as, in fact, the proof of the KKT theorem uses a rationale analogous to the one proposed here [cf. Bertsekas, 1997]).

Putting (A13) and (A14) together, we have

$$\kappa_I^H - \kappa_I^L = R \left(\frac{d}{d\beta} \gamma(\beta^{m/2})R - \frac{d}{d\beta} \gamma(\beta^{m/2+1})R \right). \quad (\text{A15})$$

Notice that, as m grows, the right-hand side goes to zero as $\frac{d}{d\beta} \gamma(\cdot)R$ is a continuous function and both $\beta^{m/2}, \beta^{m/2+1} \in [0, 2B/m]$. However, the left-hand side remains constant which leads to a contradiction. This completes the proof.